



Responsible Artificial Intelligence Program

Qisda Corporation (hereinafter referred to as “the Company”) continues to introduce artificial intelligence into internal operating processes and digital transformation applications, and promotes its Responsible Artificial Intelligence Program based on the principles of responsibility, security, transparency, and controllability. In accordance with the Artificial Intelligence Governance Policy, the Company has established an AI governance framework. Through the AI Governance Committee, information security management mechanisms, data governance requirements, and employee education and training, the Company strengthens risk controls across the planning, development, implementation, use, and maintenance stages of AI applications, ensuring that the use of AI tools and services complies with requirements for information security, privacy protection, transparency disclosure, and ethical use.

To mitigate potential risks arising from AI applications, including data leakage, biased outputs, misuse or abuse, and inappropriate decisions, the Company has implemented the following responsible artificial intelligence program contains the following elements:

1. Limiting access to sensitive AI capabilities (e.g. facial recognition, surveillance)
2. Distinct labeling of AI-generated content and outcomes of AI-driven decisions
3. Mechanisms to detect and correct drift or degradation of AI models over time
4. Regular assessments of deployed AI models for fairness/bias
5. Initiatives (own/with suppliers) to lower the ecological footprint of AI data centers/models



6. Appeals process for users/affected third parties to contest an AI decision or outcome
7. Quantification of the impact of AI initiatives/tools on sustainability outcomes
8. Training of employees on the ethical use and/or security of AI

Describe the detail as following:

1. Limiting access to sensitive AI capabilities (e.g., facial recognition and surveillance)

In accordance with Section 6.2 Risk Management and Classification Mechanism of the Qisda Artificial Intelligence Governance Policy, the Company has introduced a risk classification system in the design of its AI knowledge bases and internal AI applications. The system clearly defines which processes are prohibited AI applications and which are high-risk, limited-risk, or low-risk applications, each subject to different supervisory and management policies. During requirement submission, a filter is used to assist screening.

Responsible AI Governance Portal

Pursuing AI-driven operational benefits and innovation while mitigating threats to trustworthiness, confidentiality, integrity, ethical use, and availability of digital assets.

Approving
Authority
**President /
Board Level**

Executive Body
**AI Governance
Committee**

 Core Principles

 AI Lifecycle & Roles

 Risk Classifier

AI Project Risk Classifier

This tool is based on Qisda's **Risk Management and Classification Mechanism**. Check the features that apply to your AI project to determine its risk level and required compliance measures.

The following shows examples of prohibited AI applications, which are rejected by the system before entering the development stage.

Project Feature Checklist

1. Prohibited Features (Unacceptable Risk)

☒ Behavioral manipulation exploiting cognitive vulnerabilities / subliminal techniques

☐ Social credit scoring systems

☐ Unauthorized real-time biometric identification (e.g., facial recognition)

☐ Malicious exploitation of system vulnerabilities

2. Sensitive Decision & Evaluation Features (High Risk)

☐ Recruitment screening or HR evaluations

☐ Credit assessment or financial scoring

3. Interactive & Generative Features (Limited Risk)

☐ Chatbots, virtual assistants, or interactive generative AI

☐ Basic utility tools (e.g., spam filtering, code optimization)

Assessment Result

RISK LEVEL DETERMINATION

X Unacceptable Risk

Explicitly Prohibited!

Management Requirements:

Pursuant to the EU AI Act and Qisda Policy, the following practices are strictly prohibited: (1) behavioral manipulation; (2) social credit scoring; (3) unauthorized real-time biometric ID; (4) malicious exploitation. Project must be terminated immediately.

High-risk applications require independent validation and continuous monitoring of the accuracy of AI applications, as well as ongoing human control, supervision, and intervention.

● 2. Sensitive Decision & Evaluation Features (High Risk)

☐ Recruitment screening or HR evaluations

☒ Credit assessment or financial scoring

● 3. Interactive & Generative Features (Limited Risk)

☐ Chatbots, virtual assistants, or interactive generative AI

☐ Basic utility tools (e.g., spam filtering, code optimization)

Assessment Result

RISK LEVEL DETERMINATION

 **High Risk**

Strict Controls Required

Management Requirements:

1. Requires **independent validation** and continuous monitoring.
2. Must retain effective and fully documented **human-in-the-loop control**.
3. Subject to AI Governance Committee review and strict compliance with ISO 27001.

Limited-risk applications must, at a minimum, provide transparency disclosure and apply corresponding guardrails.

● 3. Interactive & Generative Features (Limited Risk)

☒ Chatbots, virtual assistants, or interactive generative AI

☐ Basic utility tools (e.g., spam filtering, code optimization)

Assessment Result

RISK LEVEL DETERMINATION

 Limited Risk

Transparency Obligations

Management Requirements:

1. **Disclosure:** Must fulfill transparency disclosure obligations (inform users they are interacting with AI).
2. **Guardrails:** Prevent unauthorized input of confidential data.



The guardrail mechanism controls both inputs and outputs involving confidential information, personal data, sensitive operating information, and inappropriate content, thereby reducing the risk that AI tools may improperly disclose sensitive information or generate inappropriate responses.

For AI applications involving high sensitivity or high risk, such as facial recognition, surveillance, unauthorized biometric identification, social credit scoring, behavioral manipulation, or systems that exploit vulnerabilities, the Company applies controls based on risk classification principles and prohibits the use or deployment of AI applications with unacceptable risk. Such AI applications are expressly prohibited from development.

High-risk AI applications involving sensitive data or important decisions shall be implemented in accordance with the Company's internal information security, data governance, and authorization management requirements, ensuring that the AI usage scope, access rights, and data sources are subject to controlled management. When permissions for AI applications are requested, forms are used to clearly control authorization and restrict functional access.

The following example shows that each AI assistant application must restrict the files it can access and the authorized users or departments.



 Gemini 3.5 Flash

Prompt ✨

You are an expert in circuit testing. Using the Excel file I provide for your test items, please cross-reference each test item with the provided PDF circuit diagram one by one. If it passes, write PASS; if it fails, write FAIL; if it cannot be determined, please fill in NA. For every test, please specify the page number where you found the answer and provide an explanation of the answer.

Advanced Settings >

Default question ⓘ

You can ask in this pattern

Add

Files & Collections

2 files selected

Show the files in the chat ☐

Share with everyone ☐

Share with department ⓘ

0 department selected

Share with user ⓘ

0 user selected

Collaborate with user ⓘ

0 user selected

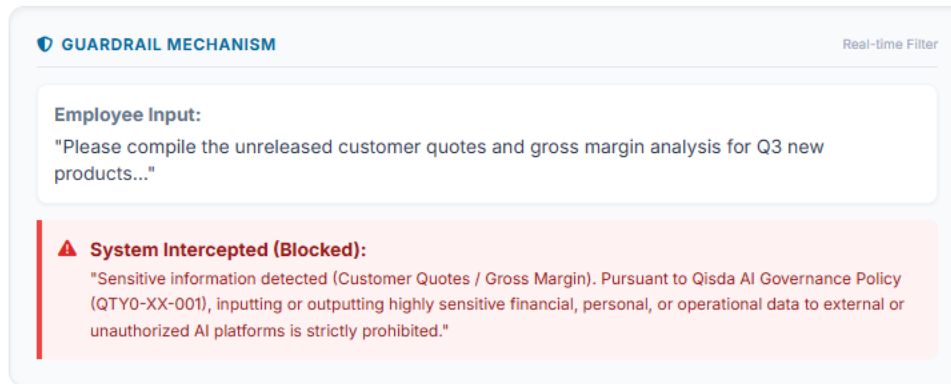
Delete

Cancel

Save



The following example shows the use of guardrail engineering mechanisms on the Company's internal generative AI platform to restrict the output of sensitive information. This helps employees reduce the risk of AI tools improperly disclosing sensitive information or generating inappropriate responses when using the platform, and avoids impacts on relevant rights and interests.



2. Distinct labeling of AI-generated content and outcomes of AI-driven decisions

In accordance with Section 6.3.2 Transparency and Explainability of the Qisda Artificial Intelligence Governance Policy, appropriate disclosure shall be made when AI is used to assist business activities, so that stakeholders do not mistake AI outputs for human-reviewed or formal decision outcomes. On all Company-authorized AI platforms, AI



assistants clearly identify themselves as AI assistants before responding to users, enabling users to understand that the response is generated or assisted by an AI tool.

For AI outputs involving user rights and interests or potentially affecting business judgment, users shall still conduct necessary checks, human review, and professional judgment in accordance with the Company's internal processes, and shall not treat AI outputs directly as the basis for final decisions.

The following actual case shows that QChatAI identifies itself as an AI assistant before providing a response.

QCHATAI RESPONSE EXAMPLE (TRANSPARENCY)



QChatAI Assistant :

"Hello! I am your **AI Assistant**. Regarding your inquiry on the ISO 27001 standards, here is the compiled information and compliance checklist for your reference..."

(Note: This report is AI-assisted and generated for decision support. The final decision and business judgment must be subject to human review.)

 Meets ****QTY0-XX-001 Transparency Principle****: Users must be explicitly informed when interacting with an AI.

In addition, AI-generated reports also indicate that AI only supports decision-making.



[2026-06-26] ABS Plastics Supply Chain Strategic Analysis Report

no_reply@qisda.com

To: recipient@qisda.com | Cc: management@qisda.com

- [96] SunSirs: Domestic Butadiene Enterprise Auction and Sales Status on June 23 (SunSirs)
- [97] SunSirs: Asian Butadiene Market Closing Status on June 22 (SunSirs)
- [98] SunSirs: International Butadiene Market Closing Status on June 22 (SunSirs)
- [99] June 23 SunSirs Butadiene Benchmark Price: 9,633.33 RMB/ton (SunSirs)
- [100] SunSirs Butadiene Average Spread Continued to Narrow Negatively to -659.00 RMB/ton on June 22 (SunSirs)
- [101] Butadiene Commodity Price Dynamics (2026-06-22) (SunSirs)
- [102] PriceSeek Alert: Domestic Butadiene Market Analysis on June 22 (SunSirs)
- [103] SunSirs: Domestic Butadiene Enterprise Auction and Sales Status on June 22 (SunSirs)
- [104] PriceSeek Alert: Butadiene Factory Price Cut, Short-term Market Weak and Volatile (SunSirs)
- [105] SunSirs: Sinopec Butadiene Factory Price on June 22 (SunSirs)
- [106] June 22 SunSirs Butadiene Benchmark Price: 9,806.67 RMB/ton (SunSirs)
- [107] SunSirs Butadiene Average Spread was -682.33 RMB/ton on June 21, Shifting from Negative Expansion to Narrowing (SunSirs)



AI Analysis Disclaimer & Verification Notice:

Data for Reference Only: This report is compiled by AI through retrieving public information. The analysis does not constitute any investment advice.

Source Verification Obligation: AI may produce errors due to differences in webpage structure parsing. Please click the original URL links in [References] for secondary verification to ensure data and date accuracy.

Timeliness Reminder: The market is highly volatile. For real-time fluctuations after the report is generated, please refer to official real-time quotes.

2026-06-26 09:05 (UTC+8 Asia/Taipei) · gemini-3.5-flash



3. Mechanisms to detect and correct drift or degradation of AI models over time

To ensure that AI systems remain stable, reliable, and aligned with expected performance during continuous use, the Company has established a regular review and continuous improvement mechanism to monitor and evaluate the performance of the AI models and related applications it uses.

In principle, the Company conducts regular reviews at least semi-annually. Using consistent test data, evaluation metrics, and actual user feedback, it examines AI system performance in terms of accuracy, relevance, reliability, security, and other applicable indicators, and compares the results against established baselines or prior results to identify potential model drift, performance degradation, or other abnormalities.

If the evaluation results indicate a significant decline in AI system performance, deviation from expected outcomes, or potential risk, the Company will adopt appropriate improvement measures based on the actual cause, including adjusting the model or system settings, optimizing prompts and knowledge sources, updating data or security control mechanisms, replacing the model or suspending the relevant application where necessary, and conducting re-validation after improvement.

In addition, for major model version changes, data source changes, system architecture adjustments, or other material changes that may affect AI performance, the Company may conduct additional assessments outside the regular review cycle based on the level of risk.

Through a continuous management process of regular evaluation, abnormality identification, improvement, and re-validation, the Company is committed to reducing the risks caused by AI model drift or degradation over time and ensuring that AI systems continue to meet their intended purposes and Responsible AI governance requirements.



4. Regular assessments of deployed AI models for fairness/bias

In accordance with Section 6.3.4 Fairness and Non-Discrimination of the Qisda Artificial Intelligence Governance Policy, for the three major commercial platforms introduced internally (Azure OpenAI, AWS Bedrock Claude, and Google Gemini), the Company implements the principle of fairness and non-discrimination through the following three-layer safeguard mechanism and conducts automated evaluations on a semi-annual basis:

1. Input and RAG Data Auditing

Pursuant to the Qisda AI Governance Policy, all data imported into RAG knowledge bases, such as RD Lesson Learnt and HR knowledge bases, must undergo de-biasing and technical accuracy review before vectorization to ensure the diversity and compliance of reference data.

System prompt constraints: Neutrality and non-discrimination instructions are embedded in system prompts such as QChatAI, strictly prohibiting unreasonable differential treatment based on sensitive attributes such as gender, age, or department.

2. Enable native safety and fairness guardrails of the three major platforms

Azure AI Content Safety: Custom filters are configured to block hate speech and bias related to sensitive attributes.



AWS Bedrock Guardrails: Sensitive information masking (PII Masking) and topic-specific refusal responses are configured.

Gemini Safety Settings: The protection levels for harassment and discriminatory speech are set to the highest level.

3. Semi-Annual Evaluation and Human-in-the-Loop Review

Semi-annual automated evaluation: Automated evaluation tools, such as Azure Prompt Flow, are used semi-annually to conduct bias scoring and fairness audits against a pre-established Golden Dataset.

Human-in-the-Loop supervision and intervention: Final decisions for high-risk applications must retain human review and shutdown authority.

Appeal and improvement mechanism: If employees identify biased AI responses, they may file an immediate appeal through the Feedback mechanism. The Company may collect information on the accuracy, appropriateness, and user experience of AI responses as a reference for subsequent optimization of knowledge base content, prompt design adjustments, response quality improvement, and strengthened risk controls, with backend data processing personnel conducting the optimization. For major bias-related issues, appeals may also be submitted at any time to the Company's CSR mailbox, csr@Qisda.com, for matters related to corporate sustainability, ESG, and other relevant topics, and the AI Governance Committee will conduct investigation and correction.

The Company also requires AI systems to retain necessary records to facilitate the tracking of data sources, inputs and outputs, key settings, and human review records.



The following case shows that employees may rate AI responses through the Feedback mechanism, providing a basis for subsequent AI improvement.

QChatAI Response Evaluation

Help us improve AI safety and accuracy

Are you satisfied with this AI response?

COMMENTS / REPORT BIAS OR ERRORS:

Please enter your specific suggestions, report biases, or file an appeal...

Submit Feedback & Appeal

QISDA AI GOVERNANCE POLICY

Feedback & Appeal Mechanism

If employees identify any bias, inaccuracies, or inappropriate content in AI responses, they can immediately submit an appeal through the Feedback mechanism. This allows the company to collect data on response accuracy, appropriateness, and user experience.

1

Data Collection & Analysis

Metrics on accuracy and appropriateness are gathered to identify system vulnerabilities.

2

Backend Optimization

Backend data processors optimize RAG knowledge bases (e.g., RD Lesson Learnt, HR Knowledge Base) and adjust prompt designs.

3

Risk Mitigation

Improves system response quality and strengthens overall risk control boundaries.

The following is the framework diagram for bias assessment.

Multi-Cloud AI Bias Mitigation & Fairness Framework

Based on: Qisda AI Governance Policy
[Ethical Compliance](#)



1. Input & RAG Data Auditing

Mitigating bias and inaccuracies at the source by auditing retrieval data and enforcing strict prompt boundaries in accordance with Qisda AI Governance Policy.



RAG Knowledge Cleansing

Audit internal **RD Lesson Learnt** and **HR Knowledge Base** documents to eliminate historical biases and technical inaccuracies before vectorization.

System Prompt Constraints

Embed mandatory ethical guidelines in system prompts to forbid differential treatment based on sensitive attributes.

**Aligns with Data Minimization & Purpose Limitation.*

2. Active Cloud Gateways

Leveraging native enterprise safety tools across Azure, AWS, and Google Cloud to block biased outputs in real-time.



Azure Content Safety

Hate Speech Filters



AWS Bedrock Guardrails

PII & Denial Topics



Gemini Safety Settings

Harassment Thresholds



Real-time Interception: Automatically blocks or sanitizes outputs that violate Qisda's non-discrimination thresholds.



3. Evaluation & Oversight

Systematic auditing using golden datasets combined with human override and feedback mechanisms.

A

Semi-Annual Auditing

Run automated benchmark tests **every six months** against a diverse dataset to detect demographic bias.

B

Human-in-the-Loop (HITL)

Mandatory manual review for high-risk decisions (e.g., recruitment, credit scoring).

C

Feedback & Appeal

Enable users to flag biased responses, triggering immediate review by the AI Committee.

**Compliant with Qisda AI Governance Policy.*



5. Initiatives, whether internal or with suppliers, to lower the ecological footprint of AI data centers/models

Qisda has expressly specified in Section 6.3.7 Environmental Sustainability of the Qisda Artificial Intelligence Governance Policy that "The Company requires its own and third-party AI data centers and model service providers to strive to reduce energy consumption and carbon emissions associated with AI computing." Accordingly, when selecting language models, the Company implements the Policy across the three major commercial language model platforms through two programs: Supplier Alignment and Evidence and Green Routing. The Company confirms that the selected model computing regions are regions matched by 100% renewable energy.

1. Supplier Alignment and Evidence

The Company strictly screens and aligns with the sustainability goals of the three major cloud partners. The following are green evidence and official documentation sources from the three suppliers:

- Microsoft Azure (Azure OpenAI) Evidence:

Green commitment: Microsoft committed to matching 100% of its electricity consumption with renewable energy by 2025, and to becoming carbon negative, water positive, and zero waste by 2030.

Latest evidence (2026): Microsoft officially announced that it has achieved the milestone of matching 100% of its annual global electricity consumption with renewable energy.

Official evidence links:



[Microsoft Datacenter Sustainability](#)

[Microsoft Journey to Carbon Negative \(Official Blog\)](#)

- Amazon Web Services (AWS Bedrock) Evidence:

Green commitment: AWS committed to achieving net-zero carbon emissions by 2040 and becoming water positive by 2030, meaning that it returns more water to communities than its data centers consume. AWS accelerated its 100% renewable energy target to 2025.

Latest evidence: Through innovative infrastructure design and cooling technologies, AWS has significantly reduced embodied carbon and has achieved 75% of its 2030 water positive target.

Official evidence links:

[AWS Cloud - Amazon Sustainability](#)

[Sustainable Cloud Computing | AWS](#)

- Google Cloud (Gemini) Evidence:

Green commitment: Google committed to operating all of its data centers worldwide on 24/7 carbon-free energy (CFE) by 2030 and achieving net-zero emissions across its entire value chain.



Latest evidence: Google is the first global technology company to commit to 24/7 hourly and regional carbon-free energy matching, and it publicly discloses the CFE percentages of each cloud region.

Official evidence links:

[Google Datacenters - Operating Sustainably](#)

[Carbon Free Energy for Google Cloud Regions](#)

2. Internal Green Routing

When deploying AI models and invoking APIs, development departments give priority to data center regions that are matched by 100% renewable energy and have lower grid carbon intensity.

For example, according to the latest sustainability report officially published by Amazon, multiple AWS data center regions worldwide have achieved 100% matching of energy consumption with renewable energy procurement (100% Renewable Energy Matched). When AWS Bedrock is deployed internally at Qisda, API calls are therefore routed to regions that both support Bedrock and have achieved 100% green matching, thereby maximizing carbon footprint reduction.

US East (N. Virginia) - us-east-1: 100% green power matched and a flagship Bedrock region. It not only offers the most complete model availability and low latency, but is also supported 100% by renewable energy, reducing carbon emissions from each prompt computation at the source.



AWS official announcement and list of green regions:

[AWS Cloud – Amazon Sustainability \(Official\)](#)

(This webpage lists all AWS regions that have achieved 100% green power matching, including us-west-2, eu-central-1, ap-northeast-1, and others.)

6. Appeals process for users/affected third parties to contest an AI decision or outcome

In accordance with Section 6.3.6 Clear Boundaries and Accountability of the Qisda Artificial Intelligence Governance Policy, AI applications shall define clear boundaries for authorization scope, data use, and user responsibilities, and provide an appeal and improvement mechanism when abnormal results occur.

Qisda implements the Appeals Process for users and affected third parties through the following three stages:

- **Stage 1: Upfront Disclosure via Application Form**

AI project compliance registration: When a business unit (BU) applies to develop or introduce AI functions, it must complete the AI Project Basic Information Application Form through the internal ITARM system. In this step, the applicant unit must clearly define:

AI application type and purpose, such as conversational application, knowledge base application, or data intelligence.



Data level involved, such as sensitive personal data, confidential data, or general non-confidential data.

AI limitations and boundaries, which will be automatically converted into user-facing disclaimers.

Generation of the Transparency & Boundary Sheet: After system review and approval, the system automatically generates an "AI Transparency & Boundary Sheet" for the AI application and proactively discloses it in the user interface, enabling employees or affected third parties to clearly understand that the decision is AI-assisted and that they have the right to appeal.

- Stage 2: Active Appeal Channel via Email

When employees or affected third parties believe that an AI-generated result is biased, incorrect, or unfair, they may at any time submit an appeal to the Company's CSR mailbox, csr@Qisda.com, for matters related to corporate sustainability, ESG, and other relevant topics, and the AI Governance Committee will conduct investigation and correction.

- Stage 3: Systematic Redress and Optimization Workflow

After an appeal is submitted, the backend initiates the following four-step process:

1. Intake and logging: The appeal is immediately logged in the backend and triaged according to risk level. High-risk applications, such as those involving personal data, recruitment, or performance, trigger immediate alerts.
2. Human-in-the-Loop intervention: The system automatically assigns the project owner or business supervisor of the AI application as Human-in-Command. During the investigation, if the AI application performs automated



decision-making, the automated decision-making must be suspended and replaced by manual secondary review and verification.

3. Committee escalation: If the appeal involves systemic bias or major controversy, it will be submitted to the AI Governance Committee, composed of the General Manager, Chief Information Officer, Chief Legal Officer, and Chief Risk Officer, for final decision.
4. Redress and system optimization:

Redress: If the AI decision is confirmed to be incorrect, the AI result shall be immediately replaced with the final decision based on human review to protect the rights and interests of the party concerned.

Optimization: Backend data processing personnel conduct root cause analysis (RCA) for the case, adjust the system prompt, and optimize the RAG knowledge base, such as RD Lesson Learnt or the HR knowledge base, to prevent recurrence of the same error.

In addition to system-based channels, major bias-related issues may also be submitted at any time to the Company's CSR mailbox, csr@Qisda.com, for matters related to corporate sustainability, ESG, and other relevant topics, and the AI Governance Committee will conduct investigation and correction.

The Company also requires AI systems to retain necessary records to facilitate the tracking of data sources, inputs and outputs, key settings, and human review records.

The following case shows that employees may rate AI responses through the Feedback mechanism, providing a basis for subsequent AI improvement.



QChatAI Response Evaluation

Help us improve AI safety and accuracy

Are you satisfied with this AI response?



COMMENTS / REPORT BIAS OR ERRORS:

Please enter your specific suggestions, report biases, or file an appeal...

 Submit Feedback & Appeal

QISDA AI GOVERNANCE POLICY

Feedback & Appeal Mechanism

If employees identify any bias, inaccuracies, or inappropriate content in AI responses, they can immediately submit an appeal through the Feedback mechanism. This allows the company to collect data on response accuracy, appropriateness, and user experience.

1 Data Collection & Analysis

Metrics on accuracy and appropriateness are gathered to identify system vulnerabilities.

2 Backend Optimization

Backend data processors optimize RAG knowledge bases (e.g., RD Lesson Learnt, HR Knowledge Base) and adjust prompt designs.

3 Risk Mitigation

Improves system response quality and strengthens overall risk control boundaries.



7. Quantification of the impact of AI initiatives/tools on sustainability outcomes

Because Qisda uses three major commercial platforms (Azure, AWS, and Gemini), each API call is accompanied by precise token transmission volumes (input/output tokens), which can be directly linked to the computing energy consumption of cloud data centers. The official APIs of the three commercial platforms can be used to estimate actual carbon emissions for quantification.

Qisda's ITS Department has established a Central AI API Gateway to centrally collect invocation data from all internal AI applications, such as QChatAI and RD Lesson Learnt, and quantify indicators across two dimensions: environmental and social/efficiency:

1. Environmental: Quantification of AI Computing Carbon Footprint

To understand the environmental impact associated with generative AI services, this report primarily uses actual account-level carbon emission estimates provided by the AWS Sustainability API. This approach directly references carbon emission results calculated by AWS based on its infrastructure, regions, services, and resource usage, providing a more consistent and traceable management basis.

Based on 2025 data returned by the AWS Sustainability API, the Company's AI account recorded total annual location-based emissions of approximately 0.438443 MTCO₂e, equivalent to approximately 438 kgCO₂e. Under the market-based method, after incorporating AWS renewable energy procurement, energy attribute certificates, and other market-based energy mechanisms, annual emissions were approximately 0.096441 MTCO₂e, equivalent to approximately 96 kgCO₂e.



Further broken down by greenhouse gas inventory scope, 2025 location-based emissions can be approximately divided into Scope 1 of 0.001657 MTCO₂e, Scope 2 of 0.289745 MTCO₂e, and Scope 3 of 0.147038 MTCO₂e. Scope 2 is the main emission source, followed by Scope 3; together, the three scopes are broadly consistent with the account's total location-based emissions.

By AWS Sustainability service category, 2025 location-based emissions were approximately: Amazon EC2 0.013454 MTCO₂e, Amazon S3 0.000027 MTCO₂e, Amazon CloudFront 0 MTCO₂e, and Other 0.424962 MTCO₂e. The Other category accounted for approximately 96.9% of total account carbon emissions and was the primary emission source.

At present, AWS Sustainability does not list Amazon Bedrock as a separately queryable service category and instead includes it in the Other category. However, the Company confirms that the vast majority of spending and usage in Other services comes from Amazon Bedrock; therefore, Bedrock's annual location-based emissions may be regarded as approaching the upper-bound value of 0.425 MTCO₂e.

For AI usage management, the AI Gateway continuously records the input tokens, output tokens, model type, AWS Region, and application system for each API call. Token data is used primarily as the denominator for AI usage efficiency and carbon efficiency KPIs, rather than as the sole basis for formal carbon inventory. Internal indicators such as kgCO₂e per million tokens may subsequently be established to continuously compare resource efficiency across models, applications, and departments.

Through the above mechanism, the Company can establish an AI environmental performance management framework that uses AWS official carbon emissions data as the primary basis and token usage as a supplementary



basis. It can regularly track annual and semi-annual carbon emission changes, the proportion of emissions by service and region, and trends in AI usage and carbon efficiency. Related data may be included in the Qisda Green Computing Report as a basis for generative AI usage governance, model selection, resource optimization, and subsequent carbon reduction improvements.

2. Social & Efficiency: Quantification of Resource and Time Saved

The number of API calls directly represents the volume of tasks completed by AI for humans, which can be precisely converted into social impact and resource efficiency indicators. Each successful API call, such as automated translation, RD Lesson Learnt summarization, or code optimization, can save employees repetitive work. Accordingly, through Continuous Improvement Program (CIP) project applications and closure calculations, the Time Saved KPI can be quantified:

According to the Company's CIP statistics, the 2025 closed CIP projects involving QChatAI function development generated performance benefits totaling USD 1,152,549.



8. Training of employees on the ethical use and/or security of AI

To strengthen the Group's overall competitive advantage and help employees acquire the professional knowledge and skills required for their roles, Qisda Group continues to promote a diversified and systematic career development blueprint. The Group integrates physical and digital learning resources to build a comprehensive learning platform and provide employees with broad and in-depth learning opportunities. In alignment with business strategies and organizational development direction, the Group has built a role-based Qisda Academy framework, planned development blueprints for each academy, comprehensively cultivated employees' professional capabilities, and continuously enhanced overall talent competitiveness.

Results of the Key Employee Development Program – AI Tool Application Capability Enhancement Program:

In response to the rapid development of AI technologies and their profound impact on corporate operations and work models, the Company launched the "AI Learning New Era" development program in 2025 as an important strategy to continue and deepen existing AI talent development results. The program focuses on improving employees' practical AI tool application capabilities, promoting the practical integration of AI technologies with daily work and process improvement, and emphasizing the importance of Responsible AI, thereby further strengthening overall organizational operational efficiency, competitiveness, and Responsible AI concepts.

With the core principles of strengthening application, deepening implementation, and continuous optimization, the "AI Learning New Era" program uses systematic course design to help employees understand AI development trends, industry impacts, and practical application scenarios, and guides them to introduce AI tools into workflows and



improvement projects.

In 2025, participation in Qisda's AI-related courses reached 3,429 person-times. The courses included Responsible AI and information security-related courses, demonstrating employees' strong attention to and willingness to participate in AI learning topics. Overall course satisfaction reached 4.5 out of 5, reflecting employees' general recognition of course content and learning outcomes. Through continuous optimization of course design and teaching methods, the Company has gradually established an internal AI learning culture and promoted cross-departmental knowledge exchange and experience sharing.

At the practical application level, AI tools have been widely introduced into the Company's Continuous Improvement Program (CIP). In 2025, the usage rate of AI application projects among CIP projects reached 65%, successfully promoting process automation and operational optimization. Related results show that employees can effectively use AI technologies to support problem analysis, process simplification, and operational efficiency improvement, transforming learning outcomes into actual operational benefits.

Through the "AI Learning New Era" program, the Company continues to move toward a more professional, streamlined, and excellent work model. AI tools help employees reduce the burden of repetitive tasks, allowing them to focus on higher-value and more creative work while retaining room for continuous learning and skill improvement. In the future, the Company will continue to monitor AI technology and industry development trends, adjust course content and application directions on a rolling basis, deepen the integration of AI with operating processes, and build a solid talent foundation for sustainable corporate development and long-term competitiveness.



Benchmarking against ISO 42001: Qisda's Commitment to Strengthening AI Management Systems and Governance Maturity

Qisda has not yet applied for ISO42001 certification. However, it will continue to review and strengthen its Responsible AI Program based on AI technology development, regulatory trends, internal usage, and stakeholder expectations. Taking ISO42001 as a benchmark, the Company will subsequently evaluate the introduction of more comprehensive mechanisms for AI model drift detection, regular fairness and bias assessment, quantitative measurement of AI applications' impact on sustainability performance, and external verification of the AI management system, in order to continuously improve AI governance maturity and transparency.